



발표요약문

발표자의 요청에 따라 발표자료 대신 요약문을 제공합니다.

생성형 AI 활용과 보안

생성형 AI의 확산으로 업무 효율성은 높아졌지만 정보유출 위험이 급증했습니다. 조직 내부 문서, 코드, 전략 보고서 등이 프롬프트 입력 과정에서 외부로 노출될 수 있어 보안 통제 경계가 모호해지고 있습니다. AI는 보안 영역에서 탐지 정확도 향상, 평균 탐지·대응 시간(MTTD/MTTR) 단축 등에 활용되고 있지만, 공격자들 역시 AI로 딥페이크 피싱, 프롬프트 인젝션, 데이터 포이즈닝, 자동화된 악성코드 제작, APT 지원 등을 수행합니다. 이용자 관점에서는 데이터 유출, 출력 오용, 법적·윤리적 리스크가 있습니다. AI 서비스 제공자는 모델 파라미터, 알고리즘 유출 방지와 오남용 차단이 과제입니다. OWASP Top 10 for LLM이 제시한 주요 위협에는 프롬프트 인젝션, 모델 오염, 민감정보 유출, 시스템 프롬프트 노출 등이 있습니다. 대응 및 거버넌스 대안으로 국가사이버안보센터(NCSC)의 N2SF(국가 망 보안체계)와 생성형 AI 가이드라인 준수, AI-SPM(Security Posture Management), AI TRISM(신뢰·위험·보안관리) 도입, Private AI 구축, SWG(Security Web Gateway)로 업무용 AI 접근 통제 강화 등이 제시되고 있습니다. AI는 선택이 아닌 필수이고, 생산성과 경쟁력 향상을 위해 AI를 올바르게 이용하고 안전하게 통제하는 것은 우리 모두의 숙제입니다.